

Frontiers of Autonomous Systems

Andrea Omicini

`andreaomicini.apice.unibo.it`

Dipartimento di Informatica – Scienza e Ingegneria (DISI)
ALMA MATER STUDIORUM – Università di Bologna

Alma Graduate School
Bologna, Italy
May 8, 2014



Outline of Part I. Autonomy of Artificial Systems

- 1 Premises
- 2 On the Notion of Autonomy
- 3 Intentional Agents



Outline of Part II. System Autonomy & Self-Organisation

- 4 Autonomy in Complex Artificial Systems
- 5 Coordination for Self-Organisation & System Autonomy



Part I

Autonomy of Artificial Systems

Outline

- 1 Premises
- 2 On the Notion of Autonomy
- 3 Intentional Agents



Outline

1 Premises

- Human Imagination of Machines
 - Early Symptoms
 - Artificial Systems

2 On the Notion of Autonomy

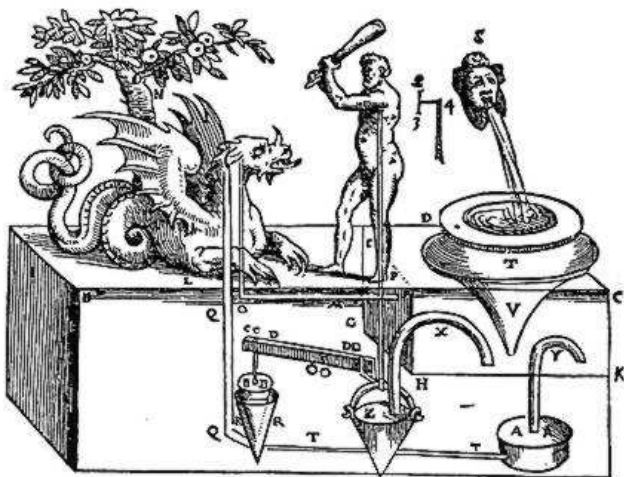
- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies



The Obsession with (Autonomous) Machines I



The Obsession with (Autonomous) Machines II

Machines doing something “by themselves”

- An obsession coming with *technique*
 - basically, representing our way to *affect* the world around us
 - possibly, according to our *goals*
- China, Greece, Italy, England, France
 - hundred years of attempts, and some success, too

The Obsession with (Autonomous) Machines III



The Obsession with (Autonomous) Machines IV



The Obsession with (Autonomous) Machines V

Aren't humans just "playing God"?

- Maybe, ok, good point.
- Fascination was so huge, too
- So strong, that *fake* automata were even quite frequent, and even famous

The Obsession with (Autonomous) Machines VI



The Obsession with (Autonomous) Machines VII

The original question

- What can human artefacts actually *do*?
- What can they *achieve*?
- What can humans *achieve* with the systems they create?

Before Autonomous Systems: Machines I

Constructing for understanding

- Building machines with
 - initiative
 - autonomy
 - knowledge

for *understanding ourselves*, and the *world* where we live

- “Playing God” to understand the world

Before Autonomous Systems: Machines II

Relieving humans from fatigue

- Goal: *substituting human* work in
 - quality
 - quantity
 - cost
- More, better, cheaper work done
 - new activities become feasible
- Which work?
 - first, *physical*
 - then, *repetitive*, enduring
 - subsequently, *intellectual*, too
 - finally, simply more *complex* for any reason—or, all reasons together

Before Autonomous Systems: Machines III

Some steps beyond

- *Delegating* human *functions* to machines
 - within already existing social structures, organisations, and processes
- Creating *new functions*
 - then, making new social structures, organisations, and processes possible
 - example: steam engines on wheels
- Essentially, changing the world we live in

Outline

1 Premises

- Human Imagination of Machines
- **Early Symptoms**
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies

Autonomous Software Creatures



Hoschie.deviantart.com

Autonomous Robot Toys?



Autonomous Vacuum Cleaners



Autonomous Lawnmowers



Autonomous Cars?



Autonomous Aircrafts?



Autonomous Weapons?



Autonomous Soccer Players & Teams???



Social Pressure

- Activities that might be *delegated* to *artificial systems* grow in *number* and *complexity*
- People & organisations are already experiencing / conceiving “somehow autonomous” systems, and ask for more
- Engineers are not yet trained on *general approaches* to build *autonomous systems*
- However, the theory of autonomous systems is at a good stage of development, and technologies are rapidly growing in terms of availability and reliability

Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- **Artificial Systems**

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies

Machines & Artificial Systems I

Systems and machines

- We call *systems* what many years ago we simply called *machines*
- Complexity has grown
- More and more we understand the many levels at which systems, their components, their mutual relationships can be described
- Furthermore, at the right level of abstraction, HW / SW systems are machines in the same acceptation as mechanical machines
- Here, we will mostly deal with two non-strictly coherent, but simple notions
 - system as a *primitive notion* (which somehow we all share to a certain extent)
 - system as an *engineer-designed entity* (“draw a line around what you call ‘a system’”)

Machines & Artificial Systems II

Artificial systems

- Here we mostly talk about *artificial systems* in general
- Systems either partially or totally designed by humans
 - either directly or indirectly
 - ? systems designed by systems?

featuring

- a *goal* (in the mind of the designer)
- a *function* (in the body of the system)

and implicitly consider the computational part as an essential one

- An artificial system, roughly speaking, is any sort of system which humans put at work by assigning it a function in order to achieve some goal

Which Sorts of Systems? I

Artificial & computational systems

- Nowadays, most (if not all) artificial systems have a prominent *computational* part
 - For this and other obvious reasons, here we focus on that sort of systems
- Computational machines
 - have both an abstract and a physical part
 - where the physical portions are often abstracted away
 - are (mostly) symbolic
 - can deal with math, logic, data, information, knowledge
 - are general-purpose machines
 - programmable, can be specialised to most purposes

Which Sorts of Systems? II

Artificial systems in context

- Most artificial systems *participate* to the activities of individuals, groups and societies
- Even more, nowadays they are mostly essential to all sorts of human activities

Socio-technical systems

- Artificial systems where both *humans* and *artificial* components play the role of system components
 - from online reservation systems to social networks
- Most of nowadays systems are just *socio-technical systems*
 - or, at least, cannot be engineering and successfully put to work without a proper socio-technical perspective in the engineering stage

Which Sorts of Systems? III

Pervasive systems

- Affecting every aspects of our everyday life
- By spreading through the whole environment where we live and act
- We live surrounded by *pervasive systems*

Situated systems

- The *physical* nature of artificial components cannot be always be forgot
- As well as the situatedness in *time* and *space*
- Along with the influence of the surrounding *environment*
- Most of the interesting systems, nowadays, are *situated systems*, too

What is the Matter with Autonomy? I

Who does what?

- This is no longer an issue
- Artificial system are very welcome to do whatever we like

Who takes the decision?

- *Autonomy* is more about *deliberation* than about *action*
- *E.g.*, for artificial weapons, the question is not “who pulls the trigger”, but rather “who decides to pull the trigger”
- Here, *ethical issues* become more than relevant

What is the Matter with Autonomy? II

Who is *responsible*?

- What is something goes wrong?
- Who is going to take responsibility – under either civil law or criminal law?
- *Legal issues* are here as least as relevant as technical issues

Outline

- 1 Premises
- 2 On the Notion of Autonomy**
- 3 Intentional Agents



Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- **Autonomy in Language**
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies

Oxford Dictionary of English (2nd Edition revised 2005)

Etimology

Early 17th cent.: from Greek *autonomia*, from *autonomos* 'having its own laws', from *autos* 'self' + *nomos* 'law'.

Dictionary

autonomy

- the right or condition of self-government
- a self-governing country or region
- freedom from external control or influence; independence.
- (in Kantian moral philosophy) the capacity of an agent to act in accordance with objective morality rather than under the influence of desires

Oxford Thesaurus of English (2nd Edition revised 2008)

Thesaurus

autonomy

- self-government, independence, self-rule, home rule, sovereignty, self-determination, freedom, autarchy;
- self-sufficiency, individualism.

Merriam-Webster I

Dictionary

autonomy

- 1 the quality or state of being self-governing; especially: the right of self-government
- 2 self-directing freedom and especially moral independence
- 3 a self-governing state

Synonyms accord, free will, choice, self-determination, volition, will

Antonyms dependence (also dependance), heteronomy, subjection, unfreedom

Merriam-Webster II

Thesaurus

autonomy

- 1 the act or power of making one's own choices or decisions:
accord, free will, choice, self-determination, volition, will
- 2 the state of being free from the control or power of another:
freedom, independence, independency, liberty,
self-determination, self-governance, self-government,
sovereignty

Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- **Autonomy in Philosophy**
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies

Internet Encyclopedia of Philosophy I

Many acceptations of autonomy

- general** an individual's capacity for self-determination or self-governance
- folk** inchoate desire for freedom in some area of one's life
- personal** the capacity to decide for oneself and pursue a course of action in one's life
- moral** the capacity to deliberate and to give oneself the moral law, rather than merely heeding the injunctions of others
- political** the property of having one's decisions respected, honored, and heeded within a political context

Internet Encyclopedia of Philosophy II

Individual autonomy

- after Kant, autonomy is an essential trait of the *individual*, and strictly related with its *morality*, represented by some high-level ethical principles
- then, with the relation between its inner self and its individual actions
- that is, mind and behaviour

Internet Encyclopedia of Philosophy III

Independence from oneself

- a more demanding notion of autonomy requires not only self-determination, but also *independence from oneself*
- this conception is connected with notions of freedom and choice, and (maybe) non-determinism
- and requires the ability of reasoning on (and possibly changing) not just one own course of actions, but one own goals

Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- **Autonomy in Military**
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies



Unmanned Systems Integrated Roadmap FY 2011-2036 I

Automatic vs. autonomous

Automatic systems are fully pre-programmed and act repeatedly and independently of external influence or control. An automatic system can be described as self-steering or self-regulating and is able to follow an externally given path while compensating for small deviations caused by external disturbances. However, the automatic system is not able to define the path according to some given goal or to choose the goal dictating its path.

Autonomous systems are self-directed toward a *goal* in that they do not require outside control, but rather are governed by laws and strategies that direct their behavior. Initially, these control algorithms are created and tested by teams of human operators and software developers. However, if *machine learning* is utilized, autonomous systems can develop modified strategies for themselves by which they select their behavior. An autonomous system is self-directed by choosing the behavior it follows to reach a human-directed goal.

Unmanned Systems Integrated Roadmap FY 2011-2036 II

Four Levels of Autonomy for Unmanned Systems (*Guy Edwards*)

Various levels of autonomy in any system guide how much and how often humans need to interact or intervene with the autonomous system:

- Human Operated
- Human Delegated
- Human Supervised
- Fully Autonomous

Unmanned Systems Integrated Roadmap FY 2011-2036 III

Level	Name	Description
1	Human Operated	A human operator makes all decisions. The system has no autonomous control of its environment although it may have information-only responses to sensed data.
2	Human Delegated	The vehicle can perform many functions independently of human control when delegated to do so. This level encompasses automatic controls, engine controls, and other low-level automation that must be activated or deactivated by human input and must act in mutual exclusion of human operation.
3	Human Supervised	The system can perform a wide variety of activities when given top-level permissions or direction by a human. Both the human and the system can initiate behaviors based on sensed data, but the system can do so only if within the scope of its currently directed tasks.
4	Fully Autonomous	The system receives goals from humans and translates them into tasks to be performed without human interaction. A human could still enter the loop in an emergency or change the goals, although in practice there may be significant time delays before human intervention occurs.

Unmanned Systems Integrated Roadmap FY 2011-2036 IV

Autonomy & Unpredictability

- The special feature of an autonomous system is its ability to be goal-directed in *unpredictable situations*.
- This ability is a significant improvement in capability compared to the capabilities of automatic systems.
- An autonomous system is able to *make a decision* based on a set of rules and/or limitations.
- It is able to determine *what information is important* in making a decision.
- It is capable of a *higher level of performance* compared to the performance of a system operating in a predetermined manner.

Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- **Autonomy in Social Sciences & AI**
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies

Autonomy as a Relational Concept [Castelfranchi, 1995] I

Autonomy as a social concept

- an agent is autonomous mostly in relation to other agents
- autonomy has no meaning for an agent in isolation

Autonomy from environment

- the Descartes' problem: (human, agent) behaviour is affected by the environment, but is not depending on the environment
- situatedness, reactivity, adaptiveness do not imply lack of autonomy

Autonomous Goals [Castelfranchi, 1995] I

Agency

- agents as teleonomic / teleologic, *goal-driven* entities
 - that is, whose *behaviour* is *not casual* under any acceptance of the term
- it might be non-deterministic, never casual

Autonomous Goals [Castelfranchi, 1995] II

Agents & Goals [Conte and Castelfranchi, 1995]

- Agents in a society can be generally conceived as either *goal-governed* or *goal-oriented* entities
 - goal-governed entities refer to the strong notion of agency, i.e. agents with some forms of cognitive capabilities, which make it possible to explicitly represent their goals, driving the selection of agent actions
 - goal-oriented entities refer to the weak notion of agency, i.e. agents whose behaviour is directly designed and programmed to achieve some goal, which is not explicitly represented
- In both cases, agent goals are *internal*

Autonomous Goals [Castelfranchi, 1995] III

Executive vs. motivational autonomy

executive autonomy — given a goal, the agent is autonomous in achieving it by itself

motivational autonomy the agent's goals are somehow self-generated, not externally imposed

Autonomy & autonomous goals

- autonomy requires *autonomous goals*
- executive autonomy is not enough for real autonomy

Goal-autonomous agent

A *goal-autonomous agent* is an agent endowed with its own *goals*

Autonomous Goals [Castelfranchi, 1995] IV

Fully socially autonomous agent

An agent is *fully socially autonomous* if

- it has its own goals: *endogenous*, not derived from other agents' will
- it is able to *make decisions* concerning *multiple conflicting goals* (being them its own goals or also goals adopted from outside)
- it adopts *goals from outside*, from other agents; it is liable to *influencing*
- it adopts other agents' goals as a consequence of a *choice* among them and other goals
- it adopts other agents' goals only if it sees the adoption as a way of enabling itself to achieve some of its own goals (i.e., the autonomous agent is a *self-interested* agent)
 - it is not possible to directly modify the agent's goals from outside: any modification of its goals must be achieved by modifying its *beliefs*
 - thus, the control over beliefs becomes a filter, an additional control over the adoption of goals
- it is impossible to change automatically the beliefs of an agent
 - the adoption of a belief is a special "decision" that the agent takes on the basis of many criteria
 - this protects its *cognitive autonomy*

Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- **Autonomy in Programming Languages**
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies

Evolution of Programming Languages: The Picture

- [Odell, 2002]

	Monolithic Programming	Modular Programming	Object-Oriented Programming	Agent Programming
Unit Behavior	Nonmodular	Modular	Modular	Modular
Unit State	External	External	Internal	Internal
Unit Invocation	External	External (CALLED)	External (message)	Internal (rules, goals)

Evolution of Programming Languages: Dimensions

Historical evolution

- Monolithic programming
- Modular programming
- Object-oriented programming
- Agent programming

Degree of modularity & encapsulation

- Unit behaviour
- Unit state
- Unit invocation

Monolithic Programming

- The basic unit of software is the whole program
- Programmer has full control
- Program's state is responsibility of the programmer
- Program invocation determined by system's operator
- Behaviour could not be invoked as a reusable unit under different circumstances
 - modularity does not apply to unit behaviour

Modular Programming

- The basic unit of software are structured loops / subroutines / procedures / ...
 - this is the era of procedures as the primary unit of decomposition
- Small units of code could actually be reused under a variety of situations
 - modularity applies to subroutine's code
- Program's state is determined by externally supplied parameters
- Program invocation determined by CALL statements and the likes

Object-Oriented Programming

- The basic unit of software are objects & classes
- Structured units of code could actually be reused under a variety of situations
- Objects have local control over variables manipulated by their own methods
 - variable state is persistent through subsequent invocations
 - object's state is encapsulated
- Object are passive—methods are invoked by external entities
 - modularity does not apply to unit invocation
 - object's control is not encapsulated

Agent-Oriented Programming

- The basic unit of software are *agents*
 - encapsulating everything, in principle
 - by simply following the pattern of the evolution
 - whatever an agent is
 - we do not need to define them now, just to understand their desired features
- Agents could in principle be reused under a variety of situations
- Agents have control over their own state
- Agents are active
 - they cannot be invoked
 - agent's control is encapsulated
- Agents are *autonomous* entities

Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- **Autonomy for Software Agents**

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies

Autonomy as the Foundation of the Definition of Agent

Lex Parsimoniae: Autonomy

- Autonomy as the only fundamental and defining feature of agents
- Let us see whether other typical agent features follow / descend from this somehow

Computational Autonomy

- Agents are autonomous as they *encapsulate* (the thread of) *control*
- Control does not pass through agent boundaries
 - only data (knowledge, information) crosses agent boundaries
- Agents have no interface, cannot be controlled, nor can they be invoked
- Looking at agents, MAS can be conceived as an aggregation of multiple distinct *loci* of control interacting with each other by exchanging information

(Autonomous) Agents (Pro-)Act

Action as the essence of agency

- The etymology of the word *agent* is from the Latin *agens*
- So, agent means “the one who acts”
- Any coherent notion of agency should naturally come equipped with a model for agent actions

Autonomous agents are pro-active

- Agents are literally active
 - Autonomous agents encapsulate control, and the rule to govern it
- Autonomous agents are pro-active by definition
- where pro-activity means “making something happen”, rather than waiting for something to happen

Agents are Situated

The model of action depends on the context

- Any “ground” model of action is strictly coupled with the context where the action takes place
- An agent comes with its own model of action
- Any agent is then strictly coupled with the environment where it lives and (inter)acts
- Agents are in this sense are intrinsically *situated*

Agents are Reactive I

Situatedness and reactivity come hand in hand

- Any model of action is strictly coupled with the context where the action takes place
- Any action model requires an adequate *representation* of the world
- Any *effective* representation of the world requires a *suitable* balance between environment *perception* and representation
- Any effective action model requires a suitable balance between environment perception and representation
 - however, any non-trivial action model requires some form of perception of the environment—so as to check action pre-conditions, or to verify the effects of actions on the environment
- Agents in this sense are supposedly *reactive* to change

Agents are Reactive II

Reactivity as a (deliberate) reduction of proactivity

- An autonomous agent could be built / choose to merely react to external events
- It may just wait for something to happen, either as a permanent attitude, or as a temporary opportunistic choice
- In this sense, autonomous agents may also be reactive

Reactivity to change

- Reactivity to (environment) change is a different notion
- This mainly comes from early AI failures, and from robotics
- It stems from agency, rather than from autonomy
- However, this issue will be even clearer when facing the issue of artifacts and environment design

(Autonomous) Agents Change the World

Action, change & environment

- Whatever the model, any model for action brings along the notion of *change*
 - an agent acts to change something around in the MAS
- Two admissible targets for change by agent action
 - agent** an agent could act to change the state of another agent
 - since agents are autonomous, and only data flow among them, the only way another agent can change their state is by providing them with some information
 - change to other agents essentially involves *communication actions*
 - environment** an agent could act to change the state of the environment
 - change to the environment requires *pragmatical actions*
 - which could be either physical or virtual depending on the nature of the environment

Autonomous Agents are Social

From autonomy to society

- From a philosophical viewpoint, autonomy only makes sense when an individual is immersed in a society
 - autonomy does not make sense for an individual in isolation
 - no individual alone could be properly said to be autonomous
- This also straightforwardly explain why any program in any sequential programming language is not an autonomous agent *per se* [Graesser, 1996, Odell, 2002]

Autonomous agents live in a MAS

- Single-agent systems do not exist in principle
- Autonomous agents live and interact within agent societies & MAS
- Roughly speaking, MAS are the only “legitimate containers” of autonomous agents

Autonomous Agents are Interactive

Interactivity follows, too

- Since agents are subsystems of a MAS, they interact within the global system
 - by essence of systems in general, rather than of MAS
- Since agents are autonomous, only data (knowledge, information) crosses agent boundaries
- Information & knowledge is exchanged between agents
 - leading to more complex patterns than message passing between objects

Autonomous Agents Do not *Need* a Goal

Agents govern MAS computation

- By encapsulating control, agents are the main forces governing and pushing computation, and determining behaviour in a MAS
- Along with control, agent should then encapsulate the *criterion* for regulating the thread(s) of control

Autonomy as self-regulation

- The term “autonomy”, at its very roots, means self-government, self-regulation, self-determination
 - “internal unit invocation” [Odell, 2002]
- This does *not* imply in any way that agents *needs* to have a goal, or a task, to be such—to be an agent, then
- However, this *does* imply that autonomy captures the cases of goal-oriented and task-oriented agents
 - where goals and tasks play the role of the criteria for governing control

Agents as Autonomous Components

Definition (Agent)

Agents are *autonomous computational entities*

genus agents are computational entities

differentia agents are autonomous, in that they encapsulate control along with a criterion to govern it

Agents are *autonomous*

- From autonomy, many other features stem
 - autonomous agents *are* interactive, social, proactive, and situated;
 - they *might* have goals or tasks, or be reactive, intelligent, mobile
 - they live within MAS, and *interact* with other agents through *communication actions*, and with the environment with *pragmatical actions*

Outline

- 1 Premises
- 2 On the Notion of Autonomy
- 3 Intentional Agents**



Intelligent Agents I

- According to a classical definition, an intelligent agent is a computational system capable of *autonomous* action and *perception* in some environment

Reminder: Computational autonomy

- Agents are autonomous as they encapsulate (the thread of) control
- Control does not pass through agent boundaries
 - only data (knowledge, information) crosses agent boundaries
- Agents have no interface, cannot be controlled, nor can they be invoked
- Looking at agents, MAS can be conceived as an aggregation of multiple distinct loci of control interacting with each other by exchanging information

Intelligent Agents II

Question: What about the other notions of autonomy?

- Autonomy with respect to other agents
- Autonomy with respect to environment
- Autonomy with respect to humans
- Autonomy with respect to oneself, moral autonomy
- ...

Question: What is intelligence to autonomy?

- Any sort of intelligence?
- Which intelligence for which autonomy?
- Which intelligent architecture for which autonomy?
- ...

Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- **Intentional Systems**
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies

Intentional Systems

- An idea is to refer to human attitudes as *intentional notions*
- When explaining human activity, it is often useful to make statements such as the following:
 - Kimi got rain tires because he believed it was going to rain
 - Nando is working hard because he wants to win his first world championship with his red car
- These statements can be read in terms of *folk psychology*, by which human behaviour can be explained and can be *predicted* through the attribution of mental attitudes, such as believing and wanting (as in the above examples), hoping, fearing, and so on.

The Intentional Stance

- The philosopher – cognitive scientist – Daniel Dennett coined the term *Intentional System* to describe entities '*whose behaviour can be predicted by the method of attributing to it belief, desires and rational acumen*' [Dennett, 1971].
- Dennett identifies several grades of intentional systems:
 - ➊ A first-order intentional system has beliefs, desires, etc.
 - Kimi believes P
 - ➋ A second-order intentional system has beliefs, desires, etc. about beliefs, desires, etc. both its own and of others
 - Kimi believes that Nando believes P
 - ➌ A third-order intentional system is then something like
 - Kimi believes that Nando believes that Felipe believes P
 - ➍ ...

The Intentional Stance in Computing

What entities can be described in terms of intentional stance?

- Human beings are prone to provide an intentional stance to almost anything
 - sacrifices for ingratiating gods' benevolence
 - animism

[McCarthy, 1979] 'Ascribing mental qualities like beliefs, intentions and wants to a machine is sometimes correct if done conservatively and is sometimes necessary to express what is known about its state [...] it is useful when the ascription helps us understand the structure of the machine, its past or future behaviour, or how to repair or improve it'

Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- **Agents with Mental States**
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies

Agents as Intentional Systems

- As computer systems become ever more complex, we need more powerful abstractions and metaphors to explain their operation
- With complexity growing, mechanistic / low-level explanations become impractical
- The intentional notions can be adopted as abstraction tools, providing us with a convenient and familiar way of describing, explaining, and predicting the behaviour of complex systems
- The idea is to use the **intentional stance** as an abstraction in computing in order to explain, understand, drive the behaviour—then, crucially, program computer systems

Agents as Intentional Systems

Strong Notion of Agency

- Early agent theorists start from a (strong) notion of agents as intentional systems
- Agents were explained in terms of mental attitudes, or mental states
- In their social abilities, agents simplest consistent description implied the intentional stance
- Agents contain an explicitly-represented – *symbolic* – model of the world (written somewhere in the working memory)
- Agents make decision on what action to take in order to achieve their goals via *symbolic* reasoning

When?

Mental states are a worth abstraction for developing agents to effectively act in a class of application domains characterized by various practical limitations and requirements [Rao and Georgeff, 1995]

- At any instant of time there are many different ways in which an environment may evolve (the environment is not deterministic)
- At any instant of time there are many actions or procedures the agent may execute (the agent is not deterministic, too)
- At any instant of time the agent may want to achieve several objectives
- The actions or procedures that (best) achieve the various objectives are dependent on the state of the environment (i.e., on the particular situation, context)
- The environment can only be sensed locally
- The rate at which computations and actions can be carried out is within reasonable bounds to the rate at which the environment evolves

Goal-Oriented & Goal-Directed Systems

- There are two main families of architectures for agents with mental states

Teleo-Reactive / Goal-Oriented Agents are based on their internal design model and their control mechanism. The goal is not figured within the internal state but it is an 'end state' for agents internal state machine

Deliberative / Goal-Directed Agents are based on a further symbolic reasoning about goals, which are explicitly represented and processed aside the control loop

Conceiving Agents with Mental States

Modelling agents based on **mental states**...

- ... eases the development of agents exhibiting complex behaviour
- ... provides us with a familiar, non-technical way of understanding and explaining agents
- ... allows the developer to build MAS by taking the perspective of a cognitive entity engaged in complex tasks – what would **I** do in the same situation?
- ... simplifies the construction, maintenance and verification of agent-based applications
- ... is useful when the agent has to communicate and interact with users or other system entities

Conceiving Agents with Mental States

- “The intentional stance is the strategy of interpreting the behaviour of an entity (person, animal, artifact, whatever) by treating it as if it were a rational agent who governed its ‘choice’ of ‘action’ by a ‘consideration’ of its ‘beliefs’ and ‘desires’ ”.
- “The scare-quotes around all these terms draw attention to the fact that some of their standard connotations may be set aside in the interests of exploiting their central features: *their role in practical reasoning*, and hence in the prediction of the behaviour of practical reasoners” [Dennett, 2007].

Agents with Mental States

Agents with Mental States

- Agents governing their behaviour on the basis of internal states that mimic cognitive mental states

Epistemic States representing agents knowledge (of their world)

- i.e., Percepts, Beliefs

Motivational States representing agents objectives

- i.e., Goals, Desires

- The process of selecting an action to execute based on the actual mental states is called *Practical Reasoning*

Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- **Intentions and Practical Reasoning**
- BDI Agents
- Agent Technologies

Practical and Epistemic Reasoning

What Practical Reasoning is

- Practical reasoning is reasoning directed towards actions—the process of figuring out what to do in order to achieve what is desired
- “Practical reasoning is a matter of weighing conflicting considerations for and against competing options, where the relevant considerations are provided by what the agent desires/values/cares about and what the agent believes.” [Bratman, 1987]

What Epistemic Reasoning is

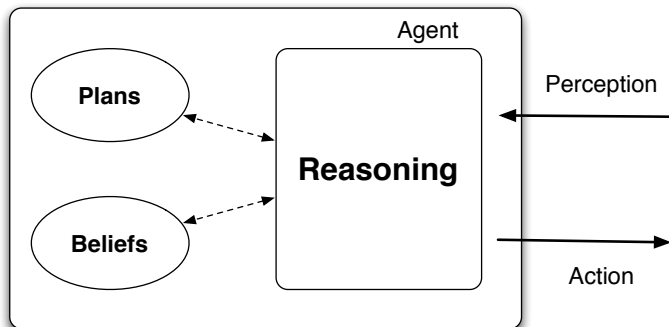
- Epistemic reasoning is reasoning directed towards knowledge—the process of updating information, replacing old information (no longer consistent with the world state) with new information

Practical Reasoning

What Practical Reasoning is

- Cognitive practical reasoning consists of two main activities
 - Deliberation** when the agent makes decision on what state of affairs the agent desire to achieve
 - Means-Ends Reasoning** when the agent makes decisions on how to achieve these state of affairs
- The output of the Deliberation phase are the *Intentions*—what agent desires to achieve, or what he desires to do
- The output of Means-Ends is in selecting a given course of actions—the workflow of actions the agent need to do to achieve the Goals

Basic Architecture of a Mentalistic Agent



Intentions and Practical Reasoning I

- ➊ Intentions pose problems for agents who need to determine ways of achieving them. If I have an intention to ϕ , you would expect me to devote resources to deciding how to bring about ϕ .
- ➋ Intentions provide a *filter* for adopting other intentions, which must not conflict. If I have an intention to ϕ , you would not expect me to adopt an intention ψ such that ϕ and ψ are mutually exclusive.
- ➌ Agents track the success of their intentions, and are inclined to try again if their attempts fail. If an agent's first attempt to achieve ϕ fails, then all other things being equal, it will try an alternative plan to achieve ϕ .
- ➍ Agents believe their intentions are possible. That is, they believe that there is at least some way that the intentions could be brought about.

Intentions and Practical Reasoning II

- 5 Agents do not believe they will not bring about their intentions. It would not be rational for me to adopt an intention to ϕ if I believed ϕ was not possible.
- 6 Under certain circumstances, agents believe they will bring about their intentions. It would not normally be rational of me to believe that I would bring my intentions about; intentions can fail. Moreover, it does not make sense that if I believe ϕ is inevitable that I would adopt it as an intention.
- 7 Agents need not intend all the expected side effects of their intentions. If I believe $\phi \rightarrow \psi$ and I intend that ϕ , I do not necessarily intend ψ also. (Intentions are not closed under implication.)
 - This last problem is known as the side effect or package deal problem. I may believe that going to the dentist involves pain, and I may also intend to go to the dentist – but this does not imply that I intend to suffer pain!

Intentions vs. Desires

- The adoption of an intention follows the rise of a given Desire (i.e. follows the adoption of a given Goal);
- Desires and Intentions are different concepts:
 - “My desire to play basketball this afternoon is merely a potential influencer of my conduct this afternoon. It must vie with my other relevant desires [...] before it is settled what I will do”.
 - “In contrast, once I intend to play basketball this afternoon, the matter is settled: I normally need not continue to weigh the pros and cons. When the afternoon arrives, I will normally just proceed to execute my intentions.” [Bratman, 1990]

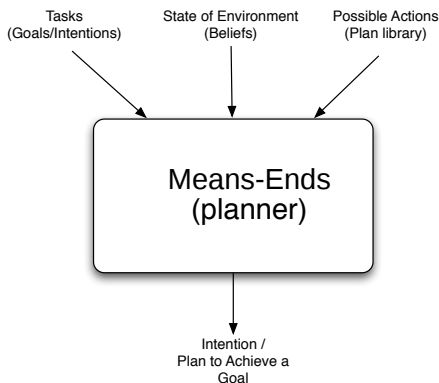
What Means-Ends Reasoning is

Means-Ends Reasoning

- The basic idea is to provide agents with three kinds of Representations:
 - Representation of Goal/Intention to achieve
 - Representation of Actions – Plans – in repertoire
 - Representation of Environment
- Given the environmental conditions, means-ends reasoning will find a Plan to achieve the adopted Goal

Mean-Ends Reasoning

- The selected intention is an emergent property, reified at runtime by selecting a given plan for achieving a given goal.



Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- **BDI Agents**
- Agent Technologies

Technical Issues for a Practical Reasoning Agent I

- Problem: Agents have Bounded Resources (bounded rationality)
 - Deliberation and Means-Ends processes are not for free: they have computational costs
 - The time taken to reason and the time taken to act are potentially unbounded, thereby destroying the reactivity and the promptness that is essential to agent survival (fitness)
- Suppose the agent
 - starts deliberating at t_0
 - begins means-ends at t_1
 - begins executing a plan at t_2
 - ends executing a plan at t_3
- Time for Deliberation is
$$t_{deliberation} = t_1 - t_0$$
- Time for Means-Ends reasoning is
$$t_{meansend} = t_2 - t_1$$

Technical Issues for a Practical Reasoning Agent II

- Agents environments are supposed to be highly dynamic
 - Many concurrent changes may occur during agent decision-making as well as during the execution of plans
- The deliberated intention is worth when it is calculated—say, at t_0
 - At time t_1 , the agent will select a goal/intention that would have been optimal if it had been achieved at t_0
- The agent runs the risk that the intention selected is no longer optimal – or no longer achievable – by the time the agent has committed to it

Technical Issues for a Practical Reasoning Agent III

So, this agent will have overall optimal behaviour in the following circumstances:

- ➊ When deliberation and means-ends reasoning take a vanishingly small amount of time
- ➋ When the world is guaranteed to remain (essentially) static while the agent is deliberating and performing means-ends reasoning, so that the assumptions upon which the choice of intention to achieve and plan to achieve the intention remain valid until the agent has completed both deliberation and means-ends reasoning
- ➌ When an intention that is optimal when achieved at t_0 (the time at which the world is observed) is guaranteed to remain optimal until t_2 (the time at which the agent has found a course of action to achieve the intention)

BDI Agents

The abstract architecture proposed by [Rao and Georgeff, 1992] comprise three dynamic and global structures representing agent's *Beliefs*, *Desires* and *Intentions* (BDI), along with an input queue of events.

- Update (write) and Query (read) operations are possible upon the three structures
 - Update operation are subject to compatibility requirements
 - Formalised constraints hold upon the mental attitudes
- The events that the system is able to recognise could be either external (i.e., coming from the environment) or internal ones (i.e., coming from some reflexive action)
 - Events are assumed to be atomic and can be recognized after they have occurred

Implementing a BDI Agent

Agent Control Loop Version 3 [Rao and Georgeff, 1995]

```
1. initialize-state();
2. while true do
3.     options := option-generator(event-queue);
4.     selected-options := deliberate(options);
5.     update-intentions(selected-options);
6.     execute();
7.     get-new-external-events();
8.     drop-successful-attitudes();
9.     drop-impossible-attitudes();
10. end-while
```

- ➊ The agent initializes the internal states
- ➋ The agent enters the main loop
- ➌ The option generator reads the event queue and returns a list of options
- ➍ The deliberator selects a subset of options to be adopted and adds these to the intention structure
- ➎ The intentions to be adopted are filtered from the selected ones
- ➏ If there is an intention to perform an atomic action at this point in time the agent executes it
- ➐ Any external events that have occurred during the interpreter cycle are then added to the event queue (the same for internal events)
- ➑ The agent modifies the intention and the desire structures by dropping successful ones
- ➒ Finally, impossible desires and intentions are dropped too

How does a BDI Agent Deliberate?

Till now, we gave no indication on how reasoning procedures of deliberation and option generation can be made sufficiently fast to satisfy the real time demands placed upon the cognitive system.

Deliberation can be decomposed in two phases:

Option Generation understand what are the available alternatives

Deliberation choose (and filter) between the adoptable goals/intentions

Chosen options are then intentions, so the agents commit to the selected ones—and executes them.

Refining Deliberation Function

Option Generation the agent generates a set of possible alternatives; represents option generation via a function, options, which takes agent's current beliefs and current intentions, and from them determines a set of options (= desires).

Deliberation the agent chooses between *competing* alternatives, and commits to the intention to achieving them. In order to select between competing options, an agent uses a *filter* function.

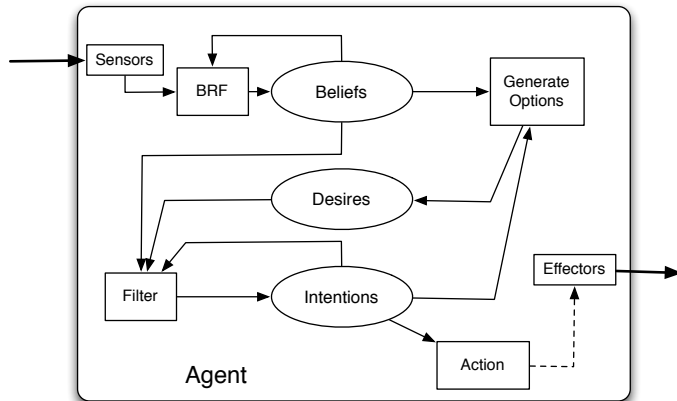
- the strategy for deliberating between goals typically is in the hands of the agent developer
- most BDI programming platforms provide mechanisms to describe under which conditions some goal should inhibit the others (Goal Formulae)
- typically, such Goal Formulae are first-order logic predicates indicating contexts and trigger conditions
- Game Theory can enter the picture: i.e., Maximizing 'Expected Utilities'

Structure of BDI Systems

BDI architectures are based on the following constructs

- ❶ A set of *beliefs*
- ❷ A set of *desires* (or *goals*)
- ❸ A set of *intentions*
 - or better, a subset of the goals with an associated stack of plans for achieving them. These are the intended actions;
- ❹ A set of *internal events*
 - elicited by a belief change (i.e., updates, addition, deletion) or by goal events (i.e. a goal achievement, or a new goal adoption).
- ❺ A set of *external events*
 - Perceptive events coming from the interaction with external entities (i.e. message arrival, signals, etc.)
- ❻ A *plan library* (repertoire of actions) as a further (static) component.

Basic Architecture of a BDI Agent [Wooldridge, 2002]



Beliefs

Beliefs

Agents knowledge is structured in Beliefs about the current state of the world

- they are informational units, typically implemented as ground sets of literals, possibly with no disjunctions or implications
- they should reflect only the information which is currently held (i.e. situated)
- they are expected to change in the future, i.e., as well as the environment changes

Plans

Plans

Plans represent the means the agent has to change the world, and to bring it closer to his desires

- they are language constructs, typically implemented in the form of procedural structures
- plans have a 'body', describing the workflow of activities (actions) that have to be executed for plan execution to be successful
- the conditions under which a plan can be chosen as an option are specified in an *invocation condition* (triggering event) and a *pre- or context- condition* (situation that must hold for the plan to be executable)

Intentions

Intentions

Intentions are emergent properties reified at runtime by selecting a given plan for achieving a given goal

- Represented 'on-line' using a run-time stack of hierarchically plans related to the ongoing adopted goals
- Similarly to how Prolog interpreter handle clauses
- Multiple intention stacks can coexist, either running in parallel, suspended until some condition occurs, or ordered for execution in some way

Outline

1 Premises

- Human Imagination of Machines
- Early Symptoms
- Artificial Systems

2 On the Notion of Autonomy

- Autonomy in Language
- Autonomy in Philosophy
- Autonomy in Military
- Autonomy in Social Sciences & AI
- Autonomy in Programming Languages
- Autonomy for Software Agents

3 Intentional Agents

- Intentional Systems
- Agents with Mental States
- Intentions and Practical Reasoning
- BDI Agents
- Agent Technologies

BDI Agents Programming Platforms

JASON (Brasil) <http://jason.sourceforge.net/> Agent platform and language for BDI agents based on AgentSpeak(L)

JADEX (Germany) <http://jadex.informatik.uni-hamburg.de/> Agent platform for BDI and Goal-Directed Agents

2APL (Netherlands) <http://www.cs.uu.nl/2apl/> Agent platform providing programming constructs to implement cognitive agents based on the BDI architecture

3APL (Netherlands) <http://www.cs.uu.nl/3apl/> A programming language for implementing cognitive agents

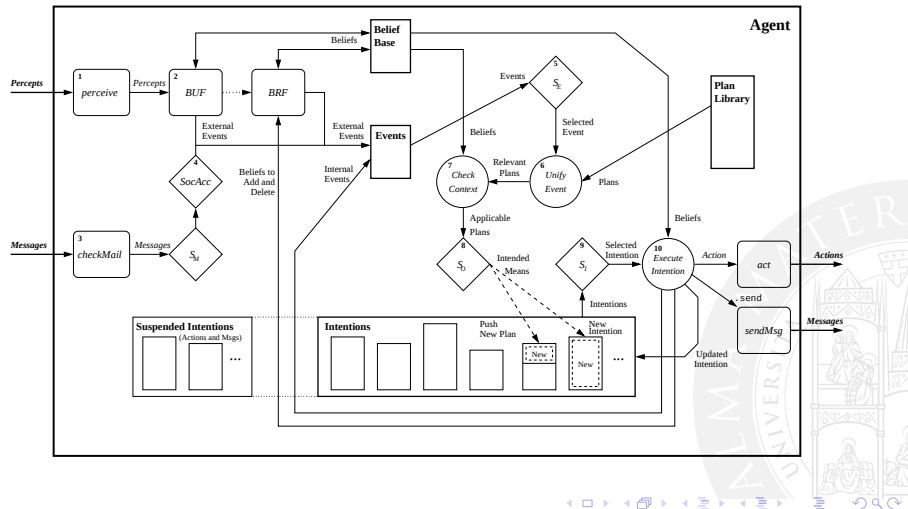
PRACTIONIST (Italy) <http://www.practitioner.org/> Framework built on the Bratman's theory of practical reasoning to support the development of BDI agents

Jason [Bordini et al., 2007]

- Developed by Jomi F. Hübner and Rafael H. Bordini
- Jason implements the operational semantics of a variant of AgentSpeak [Bordini and Hübner, 2006]
- Extends AgentSpeak, which is meant to be the language for defining agents
- Adds a set of powerful mechanism to improve agent abilities
- Extensions aimed at a more practical programming language
 - High level language to define agents (goal oriented) behaviour
 - Java as low level language to realise mechanisms (i.e. agent internal functions) and customise the architecture
- Comes with a framework for developing multi-agent systems¹

¹<http://jason.sourceforge.net/>

Jason Architecture



Jason Reasoning Cycle

- ➊ Perceiving the environment
- ➋ Updating the belief base
- ➌ Receiving communication from other agents
- ➍ Selecting 'socially acceptable' messages
- ➎ Selecting an event
- ➏ Retrieving all relevant plans
- ➐ Determining the applicable plans
- ➑ Selecting one applicable plan
- ➒ Selecting an intention for further execution
- ➓ Executing one step of an intention

Part II

System Autonomy & Self-Organisation



Outline

- 4 Autonomy in Complex Artificial Systems
- 5 Coordination for Self-Organisation & System Autonomy

Complex Systems

... by a complex system I mean one made up of a large number of parts that interact in a non simple way [Simon, 1962]

Which “parts” for complex systems?

- is autonomy of “parts” a necessary precondition?
- is it also sufficient?

Which kind of systems are we looking for?

- what is *autonomy* for a system as a whole?
- where could we find significant examples?

Nature-inspired Models

Complex natural systems

- such as physical, chemical, biochemical, biological, social systems
- natural system exhibit *features*
 - such as distribution, openness, situation, fault tolerance, robustness, adaptiveness, ...
- which we would like to understand, capture, then bring to *computational* systems

Nature-Inspired Computing (NIC)

- For instance, NIC [Liu and Tsui, 2006] summarises decades of research activities
- putting emphasis on
 - *autonomy* of components
 - *self-organisation* of systems

Outline

4 Autonomy in Complex Artificial Systems

• Self-Organisation

- Self-Organisation and Emergence in Natural Systems
- Self-Organisation and Emergence in Artificial Systems
- Self-Organising Systems & Autonomy
- Multi-Agent Systems vs. Self-Organising Systems

5 Coordination for Self-Organisation & System Autonomy

- Coordination Models
- Linda & Tuple-based Coordination
- Nature-inspired Coordination
- Tuple-based models for Nature-inspired Coordination
- Coordination Technologies

Intuitive Idea of Self-Organisation

- Self-organisation generally refers to the internal process leading to an increasing level of organisation
- *Organisation* stands for relations between parts in term of structure and interactions
- *Self* means that the driving force must be internal, specifically, distributed among components

History of Self-Organisation

- The idea of the spontaneous creation of organisation can be traced back to René Descartes
- According to the literature, the first occurrence of the term Self-Organisation is due to a 1947 paper by W. Ross Ashby [Ashby, 1947]
- Ashby defined a system to be self-organising if it changed its own organisation, rather being changed from an external entity

Elements of Self-Organisation

Increasing order — due to the increasing organisation

Autonomy — interaction with external world is allowed as long as the control is not delegated

Adaptive — suitably responds to external changes

Dynamic — it is a process not a final state

Self-Organisation in Sciences

- Initially ignored, the concept of self-organisation is present in almost every science of complexity, including
 - Physics
 - Chemistry
 - Biology and Ecology
 - Economics
 - Artificial Intelligence
 - Computer Science



History of Emergence

- Emergence is generally referred as the phenomenon involving global behaviours arising from local components interactions
- Although the origin of the term emergence can be traced back to Greeks, the modern meaning is due to the English philosopher G.H. Lewes (1875)
- With respect to chemical reactions, Lewes distinguished between *resultants* and *emergents*
 - Resultants are characterised only by their components, i.e. they are reducible
 - Conversely, emergents cannot be described in terms of their components

Definition of Emergence

- We adopt the definition of emergence provided in [Goldstein, 1999]

Emergence [..] refers to the arising of novel and coherent structures, patterns, and properties during the process of self-organisation in complex systems. Emergent phenomena are conceptualised as occurring on the macro level, in contrast to the micro-level components and processes out of which they arise.

Emergence vs. Holism

- Emergence is often, and imprecisely, explained resorting to *holism*
- Holism is a theory summarisable by the sentence *the whole is more than the sum of the parts*
- While it is true that an emergent pattern cannot be reduced to the behaviour of the individual components, emergence is a more comprehensive concept

Properties of Emergent Phenomena

Novelty — unpredictability from low-level components

Coherence — a sense of identity maintained over time

Macro-level — emergence happens at an higher-level w.r.t. to components

Dynamism — arise over time, not pre-given

Ostensive — recognised by its manifestation

Requirements for Emergency

Emergence can be exhibited by systems meeting the following requirements

Non-linearity — interactions should be non-linear and are typically represented as feedback-loops

Self-organisation — the ability to self-regulate and adapt the behaviour

Beyond equilibrium — non interested in a final state but on system dynamics

Attractors — dynamically stable working state

Definition of Self-Organisation

Widespread definition of Self-Organisation by [Camazine et al., 2001]

Self-organisation is a process in which pattern at the global level of a system emerges solely from numerous interactions among the lower-level components of the system. Moreover, the rules specifying interactions among the system's components are executed using only local information, without reference to the global pattern.

- It is evident that the authors conceive self-organisation as the source of emergence
- This tendency of combining emergence and self-organisation is quite common in biological sciences
- In the literature there is plenty of misleading definitions of self-organisation and emergence [De Wolf and Holvoet, 2005]

Outline

4 Autonomy in Complex Artificial Systems

- Self-Organisation
- **Self-Organisation and Emergence in Natural Systems**
- Self-Organisation and Emergence in Artificial Systems
- Self-Organising Systems & Autonomy
- Multi-Agent Systems vs. Self-Organising Systems

5 Coordination for Self-Organisation & System Autonomy

- Coordination Models
- Linda & Tuple-based Coordination
- Nature-inspired Coordination
- Tuple-based models for Nature-inspired Coordination
- Coordination Technologies

Natural Systems

- Natural systems can be broadly thought as [Di Marzo Serugendo et al., 2011]
 - Physical systems
 - Biological systems
 - Social systems

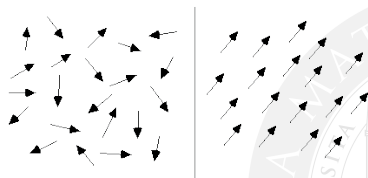
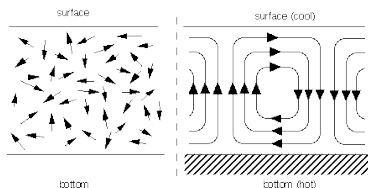


Physics and Chemistry

- Theory of self-organisation were originally developed within Physics and Chemistry
- Most typical features included
 - when the system reaches a *critical threshold*, an immediate change occurs
 - self-organisation can be observed *globally*

Self-Organisation of Matter

- Self-organisation of matter happens in several fashion
- In magnetisation, spins spontaneously align themselves in order to repel each other, producing an overall strong field
- Bénard cells is a phenomena of convection where molecules arrange themselves in regular patterns because of the temperature gradient

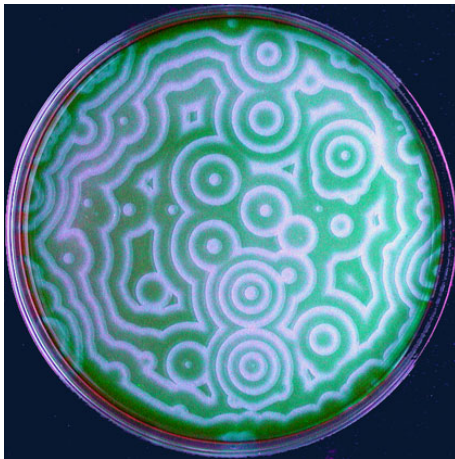


*The left hand side picture displays Bénard cells.
The right hand side picture displays magnetisation.*

Belousov-Zhabotinsky Reaction I

- Discovered by Belousov in the 1950s and later refined by Zhabontinsky, BZ reactions are a typical example of far-from-equilibrium system
- Mixing chemical reactants in proper quantities, the solution colour or patterns tend to oscillate
- These solutions are referred as *chemical oscillators*
- There have been discovered several reactions behaving as oscillators

Belousov-Zhabotinsky Reaction II



A snapshot of the Belousov-Zhabotinsky reaction.

Living Organisms

- Self-organisation is a common phenomenon in subsystems of living organisms
- An important field in biological research is the determination of invariants in the evolution of living organisms
 - in particular the spontaneous appearance of order in living complex systems due to self-organisation
- In biological research, self-organisation essentially means the global emergence of a particular behaviour or feature that cannot be reduced to the properties of individual system's components—such as molecules and cells

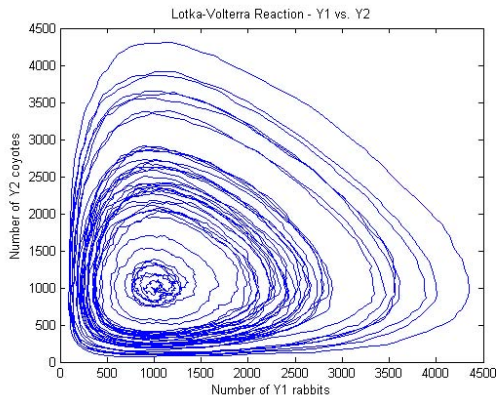
Prey-Predator Systems

- The evolution of a prey-predator systems leads to interesting dynamics
- These dynamics have been encoded in the Lotka-Volterra equation [Solé and Bascompte, 2006]
- Depending on the parameters values the system may evolve either to overpopulation, extinction or periodical evolution
- The Lotka-Volterra equation:

$$\frac{dx}{dt} = x(\alpha - \beta y)$$

$$\frac{dy}{dt} = -y(\gamma - \delta x)$$

Lotka-Volterra Equation



A chart depicting the state space defined by the Lotka-Volterra equation.

Synchronised Flashing in Fireflies I

- Some species of fireflies have been reported of being able to synchronise their flashing [Camazine et al., 2001]
- Synchronous flashing is produced by male during mating
- This synchronisation behaviour is reproducible using simple rules
 - Start counting cyclically
 - When perceive a flash, flash and restart counting

Synchronised Flashing in Fireflies II



A photo of fireflies flashing synchronously.

Schools of Fishes



School of fishes exhibit coordinated swimming: this behaviour can be simulated based on speed, orientation, and distance perception [Camazine et al., 2001].

Flocks of Birds



The picture displays a flock of geese: this behaviour can be simulated based on speed, orientation, and distance perception [Camazine et al., 2001].

Insects Colonies

- Behaviours displayed by social insects have always puzzled entomologist
- Behaviours such as nest building, sorting, routing were considered requiring elaborated skills
- For instance, termites and ants build very complex nests, whose building criteria are far than trivial, such as inner temperature, humidity, and oxygen concentration

Termites Nest in South Africa



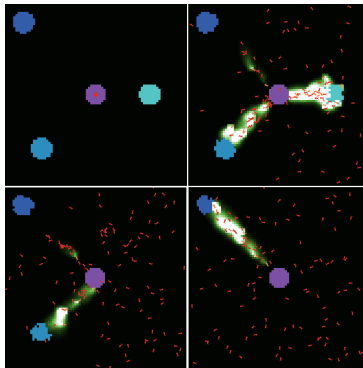
The picture displays the *Macrotermes michealseni* termite mound of southern Africa.

Trail Formation in Ant Colonies



The picture food foraging ants. When carrying food, ants lay pheromone, adaptively establishing a path between food source and the nest. When sensing pheromone, ants follow the trail to reach the food source.

Simulating Food Foraging



The snapshots display a simulation of food foraging ants featuring a nest and three food sources. Ants find the shortest path to each source and consume first the closer sources. When no longer reinforced, the pheromone eventually evaporates.

Outline

4 Autonomy in Complex Artificial Systems

- Self-Organisation
- Self-Organisation and Emergence in Natural Systems
- **Self-Organisation and Emergence in Artificial Systems**
- Self-Organising Systems & Autonomy
- Multi-Agent Systems vs. Self-Organising Systems

5 Coordination for Self-Organisation & System Autonomy

- Coordination Models
- Linda & Tuple-based Coordination
- Nature-inspired Coordination
- Tuple-based models for Nature-inspired Coordination
- Coordination Technologies

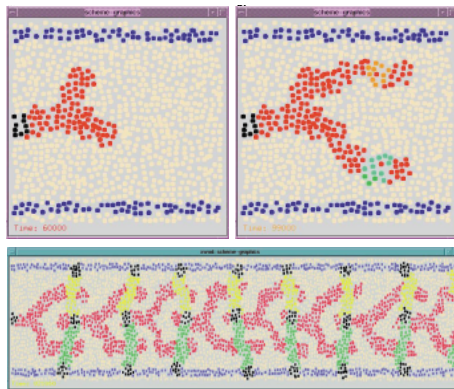
Swarm Intelligence

- *Swarm intelligence* is a problem solving approach inspired by collective behaviours displayed by social insects [Bonabeau et al., 1999, Bonabeau and Théraulaz, 2000]
- It is not a uniform theory, rather a collection of mechanisms found in natural systems having applications to artificial systems
- Applications of Swarm Intelligence include a variety of problems such as task allocation, routing, synchronisation, sorting
- In Swarm Intelligence, the most successful initiative is Ant Colony Optimisation

ACO: Ant Colony Optimisation

- ACO [Dorigo and Stützle, 2004] is a population-based metaheuristic that can be used to find approximate solutions to difficult optimisation problems
- A set of software agents called artificial ants search for good solutions to a given optimisation problem
- To apply ACO, the optimisation problem is transformed into the problem of finding the best path on a weighted graph
- ACO provided solutions to problems such as VRP-Vehicle Routing Problem, TSP-Travelling Salesman Problem and packet routing in telecommunication networks

Amorphous Computing

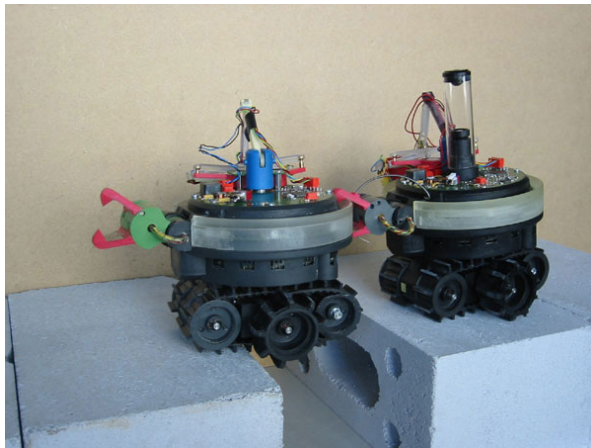


An amorphous computing [Abelson et al., 2000] medium is a system of irregularly placed, asynchronous, locally interacting identical computing elements

Autonomic Computing

- An industry driven research field initiated by IBM [Kephart and Chess, 2003], mostly motivated by increasing costs in systems maintenance
- Basic idea: applying self-organising mechanisms found in human nervous system to develop more robust and adaptive systems
- Applications range from a variety of problems such as power saving, security, load balancing

SWARM-BOTS



SWARM-BOTS [Dorigo et al., 2005] was a project funded by European Community tailored to the study of self-organisation and self-assembly of modular robots

AGV – Automated Guided Vehicles

- Stigmergy has been successfully applied to several deployments of Automated Guided Vehicles [Weyns et al., 2005, Sauter et al., 2005]
- Basically, the AGVs are driven by digital pheromones fields in the same way ants perform food-foraging



Pictures of AGVs

Outline

4 Autonomy in Complex Artificial Systems

- Self-Organisation
- Self-Organisation and Emergence in Natural Systems
- Self-Organisation and Emergence in Artificial Systems
- **Self-Organising Systems & Autonomy**
- Multi-Agent Systems vs. Self-Organising Systems

5 Coordination for Self-Organisation & System Autonomy

- Coordination Models
- Linda & Tuple-based Coordination
- Nature-inspired Coordination
- Tuple-based models for Nature-inspired Coordination
- Coordination Technologies

Are SOS Autonomous?

- They are *adaptive*, in that they properly respond to external stimuli
- So their autonomy from the environment is *partial*
- At the same time, they are *self-governed*, in that their evolution is self-driven, in some essential sense—it is at least teleonomic
- So, their autonomy is evident, as well

Systems as Agents?

- In the following, we take as understood the fact that the notion of *autonomy* applies to *systems*
 - Our implicit assumption is that *users* (generally) and *designers* (at some point) consider a *system as a whole*, and conceive it as such
 - that is, as a computational system with its own computational autonomy—which for us means an agent, at a certain level of abstraction
- this basically means that we can evaluate other notions of autonomy for a system as a whole

How Much Autonomy?

- Good design of a SOS provides the goals to be achieved, and the means to self-organise the system structure accordingly
- How much autonomy in that?
- How much autonomy from the designer, from the user, from the environment, overall?

Which Autonomy for SOS?

- Self-organising systems (SOS) exhibit some autonomy by definition
 - their evolution over time is not pre-defined by the designer
 - in this sense, SOS are autonomous with respect to the designer
 - however, any evolution of a well-engineered SOS tends towards the tasks / goals assigned by the designer
 - in this sense, SOS are not autonomous with respect to the designer
 - their evolution over time is not influenced by the environment, but is not directly driven by it
 - in this sense, SOS are autonomous with respect to the environment
- Most of the SOS we know are natural systems, where it is not clear whether one can say that the goals are somehow self-generated
- However, for sure, computational SOS built from those examples are likely to show executive autonomy, without motivational autonomy

Autonomy of SOS Depends on...

- The models, mechanisms, and technologies adopted for implementing computational SOS
- ! The level of autonomy of a SOS do not depend straightforwardly on the level of autonomy of the agent components



Component vs. System Autonomy

- SOS are systems with some autonomy made of autonomous components
 - However, no clear relationship between the sort of autonomy of components and the sort of autonomy of the system can be stated a priori
- which basically means that autonomy of a SOS does not necessarily rely upon its components only
- and also means that issues like responsibility and liability require a non-trivial, non-obvious treatment

Outline

4 Autonomy in Complex Artificial Systems

- Self-Organisation
- Self-Organisation and Emergence in Natural Systems
- Self-Organisation and Emergence in Artificial Systems
- Self-Organising Systems & Autonomy
- **Multi-Agent Systems vs. Self-Organising Systems**

5 Coordination for Self-Organisation & System Autonomy

- Coordination Models
- Linda & Tuple-based Coordination
- Nature-inspired Coordination
- Tuple-based models for Nature-inspired Coordination
- Coordination Technologies

MAS 4 SOS

- Is the agent paradigm the right choice for modelling and developing SOS?
- Are agents the right abstractions for SOS components?
- Are MAS the right way to put together components of a SOS?
- In order to answer this question we have to compare requirements for SOS with features of MAS

SOS Requirements

- From our previous discussion on self-organisation and emergence, a possible basic requirements list can be given as follows:
 - *Autonomy* and *encapsulation* of behaviour
 - *Local actions* and *perceptions*
 - *Distributed environment* supporting *interactions*
 - Support for *organisation* and *cooperation* concepts

MAS Checklist

It is easy to recognise that the agent paradigm provides suitable abstractions for each aspect

- Agents for autonomy and encapsulation of behaviour
- Situated agents for local actions and perceptions
- MAS distribution of components, and MAS environment supporting interactions through *coordination*
- MAS support for organisation and cooperation concepts

Outline

4 Autonomy in Complex Artificial Systems

5 Coordination for Self-Organisation & System Autonomy

Outline

4 Autonomy in Complex Artificial Systems

- Self-Organisation
- Self-Organisation and Emergence in Natural Systems
- Self-Organisation and Emergence in Artificial Systems
- Self-Organising Systems & Autonomy
- Multi-Agent Systems vs. Self-Organising Systems

5 Coordination for Self-Organisation & System Autonomy

- Coordination Models
- Linda & Tuple-based Coordination
- Nature-inspired Coordination
- Tuple-based models for Nature-inspired Coordination
- Coordination Technologies

Interaction & Coordination

Interaction

- most of the complexity of complex computational systems – MAS included – comes from **interaction** [Omicini et al., 2006]
- along with an essential part of their expressive power [Wegner, 1997]

Coordination

- since **coordination** is essentially the science of **managing the space of interaction** [Wegner, 1997]
- **coordination models and languages** [Ciancarini, 1996] provide abstractions and technologies for the engineering of complex computational systems [Ciancarini et al., 2000]

Coordination Models for Complex Computational Systems

Coordination model as a glue

A coordination model is the glue that binds separate activities into an ensemble [Gelernter and Carriero, 1992]

Coordination model as an agent interaction framework

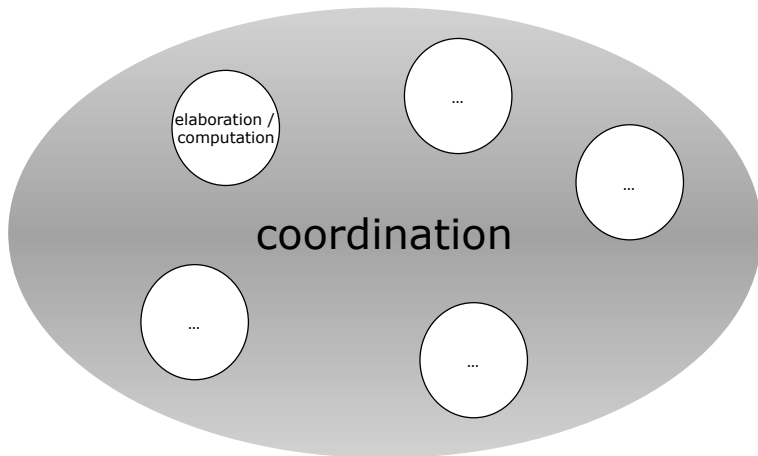
A coordination model provides a framework in which the interaction of active and independent entities called agents can be expressed [Ciancarini, 1996]

Issues for a coordination model

A coordination model should cover the issues of creation and destruction of agents, communication among agents, and spatial distribution of agents, as well as synchronization and distribution of their actions over time [Ciancarini, 1996]

What is Coordination?

Ruling the space of interaction



New Perspective on Computational Systems

Programming languages

- Interaction as an orthogonal dimension
- Languages for interaction / coordination

Software engineering

- Interaction as an independent design dimension
- Coordination patterns

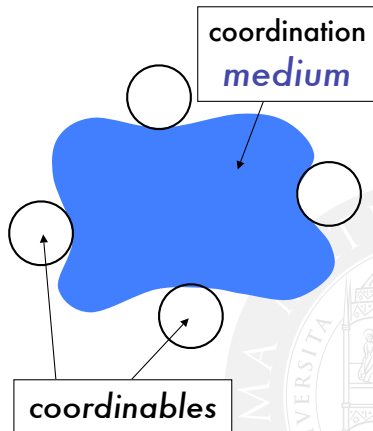
Artificial intelligence

- Interaction as a new source for intelligence
- Social intelligence

Coordination: Sketching a Meta-model

The *medium of coordination*

- “fills” the interaction space
- enables / promotes / governs the admissible / desirable / required interactions among the interacting entities
- according to some *coordination laws*
 - enacted by the behaviour of the medium
 - defining the semantics of coordination



Coordination: A Meta-model [Ciancarini, 1996]

A constructive approach

Which are the components of a coordination system?

Coordination entities Entities whose mutual interaction is ruled by the model, also called the *coordinables*

Coordination media Abstractions enabling and ruling interaction among coordinables

Coordination laws Laws ruling the observable behaviour of coordination media and coordinables, and their interaction as well

Outline

4 Autonomy in Complex Artificial Systems

- Self-Organisation
- Self-Organisation and Emergence in Natural Systems
- Self-Organisation and Emergence in Artificial Systems
- Self-Organising Systems & Autonomy
- Multi-Agent Systems vs. Self-Organising Systems

5 Coordination for Self-Organisation & System Autonomy

- Coordination Models
- **Linda & Tuple-based Coordination**
- Nature-inspired Coordination
- Tuple-based models for Nature-inspired Coordination
- Coordination Technologies

The Ancestor

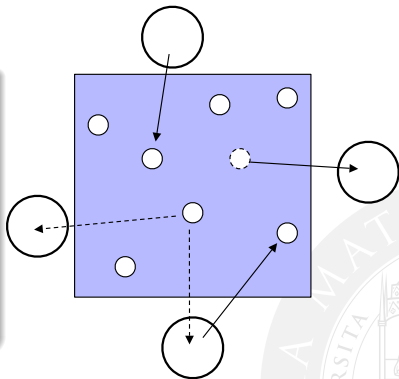
LINDA [Gelernter, 1985]

- LINDA is the ancestor of all **tuple-based coordination models** [Rossi et al., 2001]
- in LINDA, agents synchronise, cooperate, compete
 - based on **tuples**
 - available in the **tuple spaces**, working as the **coordination media**
 - by *associatively* accessing, consuming and producing tuples
- the same holds for any tuple-based coordination model

The Tuple-space Meta-model

The basics [Gelernter, 1985]

- *Coordinables* synchronise, cooperate, compete
 - based on *tuples*
 - available in the *tuple space*
 - by *associatively* accessing, consuming and producing tuples



Tuple-based / Space-based Coordination Systems

LINDA meta-model [Ciancarini, 1996]

coordination media tuple spaces

- as multiset / bag of data objects / structures called *tuples*

communication language tuples / tuple templates

tuples as ordered collections of (possibly heterogeneous) information items

templates as specifications of tuple sets

coordination language tuple space primitives

- as a set of operations to put, browse and retrieve tuples to/from the space

Multiple Tuple Spaces

ts ? out(T)

- Linda tuple space might be a bottleneck for coordination
- Many extensions have focussed on making a multiplicity of tuple spaces available to processes
 - each of them encapsulating a portion of the coordination load
 - either hosted by a single machine, or distributed across the network
- Syntax required, and dependent on particular models and implementations
 - a space for tuple space names, possibly including network location
 - operators to associate Linda operators to tuple spaces
- For instance, `ts @ node ? out(p)` may denote the invocation of operation `out(p)` over tuple space `ts` on node `node`

Main Features of Tuple-based Coordination

Main features of the Linda model

tuples a tuple is an ordered collection of knowledge chunks, possibly heterogeneous in sort

generative communication until explicitly withdrawn, the tuples generated by coordinables have an independent existence in the tuple space; a tuple is equally accessible to all the coordinables, but is bound to none

associative access templates allow accessing tuples through their content & structure, rather than by name, address, or location

suspensive semantics operations may be suspended based on unavailability of matching tuples, and be woken up when such tuples become available

Outline

4 Autonomy in Complex Artificial Systems

- Self-Organisation
- Self-Organisation and Emergence in Natural Systems
- Self-Organisation and Emergence in Artificial Systems
- Self-Organising Systems & Autonomy
- Multi-Agent Systems vs. Self-Organising Systems

5 Coordination for Self-Organisation & System Autonomy

- Coordination Models
- Linda & Tuple-based Coordination
- **Nature-inspired Coordination**
- Tuple-based models for Nature-inspired Coordination
- Coordination Technologies

Nature-inspired Coordination for MAS

Coordination issues in natural systems

- coordination issues did not first emerge in computational systems
- [Grassé, 1959] noted that in termite societies *“The coordination of tasks and the regulation of constructions are not directly dependent from the workers, but from constructions themselves.”*

Coordination as the key issue

- many well-known examples of natural systems – and, more generally, of complex systems – seemingly rely on simple yet powerful **coordination mechanisms** for their key features—such as **self-organisation**
- it makes sense to focus on **nature-inspired coordination models** as the core of complex nature-inspired MAS

Stigmergy I

Stigmergy in insect societies

- nature-inspired models of coordination are grounded in studies on the behaviour of social insects, like ants or termites
- [Grassé, 1959] introduced the notion of **stigmergy** as the fundamental coordination mechanism in termite societies
- in ant colonies, pheromones act as environment markers for specific social activities, and drive both the *individual* and the *social* behaviour of ants

Stigmergy II

Stigmergy in computational systems

- nowadays, stigmergy generally refers to a set of nature-inspired coordination mechanisms mediated by the *environment*
- *digital pheromones* [Parunak et al., 2002] and other *signs* made and sensed in a shared environment [Parunak, 2006] can be exploited for the engineering of adaptive and self-organising MAS

Chemical Coordination

Chemical reactions as (natural) coordination laws

- inspiration comes from the idea that complex physical phenomena are driven by the (relatively) simple chemical reactions
- coordinating the behaviours of a huge amount of agents, as well as the global system evolution

Chemical reactions as (computational) coordination laws

- Gamma [Banâtre and Le Métayer, 1990] is a *chemistry-inspired coordination* model—as for the CHAM (chemical abstract machine) model [Berry, 1992]
- coordination in Gamma is conceived as the evolution of a space governed by chemical-like rules, globally working as a rewriting system [Banâtre et al., 2001]

Field-based Coordination

Computational fields as coordination laws

- **field-based** coordination models like Co-Fields [Mamei and Zambonelli, 2006] are inspired by the way masses and particles move and self-organise according to gravitational/electromagnetic fields
- there, computational force fields – generated either by the mobile agents or by the pervasive coordination infrastructure – propagate across the environment, and drive the actions and motion of the agent themselves

(Bio)chemical Coordination

Chemical reactions as coordination laws

- **chemical tuple spaces** [Viroli et al., 2010] exploit the chemical metaphor at its full extent—beyond Gamma
- data, devices, and software agents are represented in terms of chemical reactants, and system behaviour is expressed by means of chemical-like laws
- which are actually **time-dependent** and **stochastic**
- embedded within the coordination medium
- **biochemical tuple spaces** [Viroli and Casadei, 2009] add *compartments*, *diffusion*, and *stochastic behaviour* of coordination primitives

Basic Issues of Nature-inspired Coordination I

Environment

Environment is essential in nature-inspired coordination

- it works as a **mediator** for agent interaction — through which agents can communicate and coordinate **indirectly**
- it is **active** — featuring autonomous dynamics, and affecting agent coordination
- it has a **structure** — requiring a notion of **locality**, and allowing agents of any sort to **move** through a **topology**

Basic Issues of Nature-inspired Coordination II

Stochastic behaviour

Complex systems typically require **probabilistic** models

- *don't know* / *don't care* **non-deterministic** mechanisms are not expressive enough to capture all the properties of complex systems such as biochemical and social systems
- probabilistic mechanisms are required to fully capture the dynamics of coordination in nature-inspired systems
- coordination models should feature (possibly simple yet) expressive mechanisms to provide coordinated systems with **stochastic behaviours**

Outline

4 Autonomy in Complex Artificial Systems

- Self-Organisation
- Self-Organisation and Emergence in Natural Systems
- Self-Organisation and Emergence in Artificial Systems
- Self-Organising Systems & Autonomy
- Multi-Agent Systems vs. Self-Organising Systems

5 Coordination for Self-Organisation & System Autonomy

- Coordination Models
- Linda & Tuple-based Coordination
- Nature-inspired Coordination
- **Tuple-based models for Nature-inspired Coordination**
- Coordination Technologies

LINDA is *not* a Nature-inspired Model

Warning

LINDA is *not* a Nature-inspired Model

So, *why* LINDA?

Why tuple-based models?

Why Tuple-based Models? I

Expressiveness

- LINDA is a sort of *core* coordination model
- making it easy to face and solve many typical problems of complex distributed systems
- *complex* coordination problems are solved with *few, simple* primitives
- whatever the model used to measure *expressiveness* of coordination, tuple-based languages are highly-expressive [Busi et al., 1998]

Why Tuple-based Models? II

Environment-based coordination

- *generative communication* [Gelernter, 1985] requires *permanent* coordination abstractions
- so, the *coordination infrastructure* provides agents with tuple spaces as *coordination services*
 - *coordination as a service* (CaaS) [Viroli and Omicini, 2006]
- they can be interpreted as *coordination artefacts* shaping computational *environment* [Omicini et al., 2004]
 - and used with different levels of awareness by both intelligent and “stupid” agents [Omicini, 2013]
- as such, they can be exploited to support *environment-based coordination* [Ricci et al., 2005]

Why Tuple-based Models? III

Extensibility

- whatever its expressiveness, LINDA was conceived as a coordination model for closed, parallel systems
- so, in fact, some relevant problems of today open, concurrent systems cannot be easily solved with LINDA either in practice or in theory
- as a result, tuple-based models have been extended with new simple yet powerful mechanisms
- generating a plethora of tuple-based coordination models
[Rossi et al., 2001]

Why Tuple-based Models? IV

Nature-inspired extensions

- LINDA may *not* be nature-inspired, but many of its extensions *are*
- many of the coordination models depicted before
 - stigmergy [Parunak, 2006]
 - field-based [Mamei and Zambonelli, 2004]
 - chemical [Viroli et al., 2010] and biochemical [Viroli and Casadei, 2009]
- along with many others, such as
 - cognitive stigmergy [Ricci et al., 2007]
 - pervasive ecosystems [Viroli et al., 2012]
- are actually **nature-inspired tuple-based coordination models**

Examples I

StoKLAIM

- StoKLAIM [De Nicola et al., 2006] – a *stochastic* extension of the LINDA-derived KLAIM model for mobile coordination [De Nicola et al., 1998] – adds distribution rates to coordination primitives—thus making it possible the modelling of non-deterministic real-life phenomena such as failure rates and inter-arrival times

SwarmLinda

- SwarmLinda [Tolksdorf and Menezes, 2004] enhances LINDA implementation with swarm intelligence to achieve features such as scalability, adaptiveness, and fault-tolerance—by modelling tuple templates as ants, featuring probabilistic behaviour when looking for matching tuples in a distributed setting

Examples II

Situated ReSpecT [Mariani and Omicini, 2013]

- ReSpecT [Omicini and Denti, 2001] generally addresses *situated dependency* by capturing *time*, *space*, and *environment events*, and supporting the definition and enforcement of *situated coordination policies*
- so, ReSpecT-programmed tuple centres can work as situated abstractions for MAS-environment coordination

Blending Metaphors

Mixing abstractions & mechanisms from different conceptual sources

- most natural systems, when observed in their whole complexity, exhibit *layers* each one featuring its own metaphors and mechanisms
- correspondingly, many novel approaches to complex MAS coordination integrate diverse sources of inspiration, e.g.:
 - TOTA [Mamei and Zambonelli, 2004] exploits mechanisms from both stigmergic and field-based coordination
 - the SAPERE coordination model for pervasive service ecosystems [Zambonelli et al., 2011, Viroli et al., 2012] integrates
 - the *chemical* metaphor for driving the evolution of coordination abstractions
 - *biochemical* abstractions for topology and diffusion
 - the notion of *ecosystem* in order to model the overall system structure and dynamics

Expressing Full Dynamics

Expressing the *full dynamics* of complex natural systems

- mostly, coordination models just capture *some* of the overall system dynamics
- which makes them basically *fail*
 - for instance, Gamma mimics chemical reactions, but does not capture essential issues in chemical processes such as reaction rates and concentration [Banâtre and Le Métayer, 1990, Banâtre et al., 2001]
 - instead, *(bio)chemical tuple spaces* fully exploit the chemical metaphor by providing time-dependent and stochastic chemical laws [Viroli et al., 2010, Viroli and Casadei, 2009]
- more generally, the goal is to allow coordinated MAS to capture and *express the full dynamics* of complex natural systems

Core Mechanisms I

Understanding the *basic elements of expressiveness*

- LINDA is a glaring example of a minimal set of coordination mechanisms providing a wide range of coordination behaviours
- the goal is understanding the minimal set of coordination primitives required to design complex stochastic behaviours
- for instance, *uniform coordination primitives* – that is, LINDA-like coordination primitives returning tuples matching a template with a uniform distribution [Gardelli et al., 2007] – seemingly capture the full-fledged dynamics of real chemical systems within the coordination abstractions

Core Mechanisms II

Issues

- Autonomy and the limits of computational systems
 - *expressiveness* of languages and technologies
 - (*technically, ethically, legally*) *admissible behaviours* of computational systems

Predicting Complex Behaviours I

Engineering unpredictable systems around predictable abstractions

- coordination models are meant to harness the complexity of complex MAS [Ciancarini et al., 2000]
- coordination abstractions are often *at the core* of complex MAS
- while this does not make complex MAS generally predictable, it makes it possible in principle to make them *partially predictable*, based on the predictability of the core coordinative behaviour
- suitably-formalised coordination abstractions, along with a suitably-defined engineering methodology, could in principle ensure the predictability of given MAS properties within generally-unpredictable MAS

Predicting Complex Behaviours II

Issues

- Autonomy and predictability
 - is unpredictability a pre-condition to choice, freedom, and so autonomy?
 - how could unpredictability coexist with well-founded notions of responsibility and liability?
- Partial predictability?

Coordination for Simulation I

Simulation of complex systems is a multidisciplinary issue

- ... ranging from physics to biology, from economics to social sciences
- no complex system of any sort can be studied nowadays without the support of suitable simulation tools
- nowadays, experiments done *in silico* are at least as relevant as those *in vitro* and *in vivo*

Coordination for Simulation II

Interaction issues are prominent in complex systems

- coordination technologies potential core of agent-based simulation frameworks
- in particular, self-organising nature-inspired coordination models are well suited for the simulation of complex systems
- so, coordination middleware could play a central role in the development of rich agent-based simulation frameworks for complex systems
- e.g., [González Pérez et al., 2013]

Issues

- Autonomy and simulation
 - what autonomy is required to simulate autonomous systems?

Knowledge-oriented Coordination I

Integrating nature-inspired with knowledge-oriented coordination

- intelligent MAS in knowledge intensive environments – as well as complex socio-technical systems, in general – require automatic understanding of data and information
- **knowledge-oriented coordination** exploits coordination abstractions enriched so as to allow for semantic interpretation by intelligent agents [Fensel, 2004, Nardini et al., 2013]
- for instance
 - chemical tuple spaces
 - SAPERE coordination abstractions and mechanisms
 - *semantic tuple centres* [Nardini et al., 2011]all relay on the semantic interpretation of coordination items

Knowledge-oriented Coordination II

Self-organisation of knowledge

- explicit search of information is going to become ineffective while the amount of available knowledge grows at incredible rates
- knowledge should autonomously organise and flow from producers to consumers
- **knowledge self-organisation** for knowledge-intensive MAS

Knowledge-oriented Coordination III

MoK (Molecules of Knowledge) [Mariani and Omicini, 2012a]

- Molecules of Knowledge is a a nature-inspired coordination model promoting **knowledge self-organisation**, where
 - sources of knowledge continuously produce and inject *atoms of knowledge* in biochemical compartments
 - knowledge atoms may then aggregate in *molecules* and diffuse
 - knowledge producers, managers and consumers are modelled as *catalysts*, whose workspaces are biochemical compartments, and their knowledge-oriented actions become enzymes influencing atoms aggregation and molecules diffusion
 - so as to make relevant knowledge spontaneously aggregate and autonomously move towards potentially interested knowledge workers
- the first application scenario for experimenting with MoK is **news management** [Mariani and Omicini, 2012b]

Knowledge-oriented Coordination IV

Issues

- Autonomy and knowledge
 - if (informed) choice is essential to freedom, and freedom is essential to autonomous choice, then *knowledge* is *essential to autonomy*
- Autonomy of knowledge?
 - what is autonomy, when *knowledge* chunks *autonomously move* towards knowledge consumers?

Outline

4 Autonomy in Complex Artificial Systems

- Self-Organisation
- Self-Organisation and Emergence in Natural Systems
- Self-Organisation and Emergence in Artificial Systems
- Self-Organising Systems & Autonomy
- Multi-Agent Systems vs. Self-Organising Systems

5 Coordination for Self-Organisation & System Autonomy

- Coordination Models
- Linda & Tuple-based Coordination
- Nature-inspired Coordination
- Tuple-based models for Nature-inspired Coordination
- Coordination Technologies

Coordination Middleware

JavaSpaces <http://www.oracle.com/technetwork/articles/javase/javaspaces-140665.html>

A Java-based high-level tool for building distributed and collaborative applications [Freeman et al., 1999]

Law-Governed Interaction (LGI) <http://www.moses.rutgers.edu>

A decentralised coordination and control mechanism for distributed systems [Minsky and Ungureanu, 2000]

TuCSoN <http://tucson.unibo.it>

A model and technology for tuple-based coordination of complex distributed systems [Omicini and Zambonelli, 1999]

Part III

Bibliography

Bibliography I



Abelson, H., Allen, D., Coore, D., Hanson, C., Homsy, G., Knight, Jr., T. F., Nagpal, R., Rauch, E., Sussman, G. J., and Weiss, R. (2000).
Amorphous computing.

Communications of the ACM, 43(5):74–82.



Ashby, W. R. (1947).

Principles of self-organizing dynamic systems.

Journal of General Psychology, 37:125–128.



Banătre, J.-P., Fradet, P., and Le Métayer, D. (2001).

Gamma and the chemical reaction model: Fifteen years after.

In Calude, C. S., Păun, G., Rozenberg, G., and Salomaa, A., editors,
*Multiset Processing. Mathematical, Computer Science, and Molecular
Computing Points of View*, volume 2235 of *LNCS*, pages 17–44.

Springer.

Bibliography II



Banâtre, J.-P. and Le Métayer, D. (1990).
The GAMMA model and its discipline of programming.
Science of Computer Programming, 15(1):55–77.



Berry, G. (1992).
The chemical abstract machine.
Theoretical Computer Science, 96(1):217–248.



Bonabeau, E., Dorigo, M., and Theraulaz, G. (1999).
Swarm Intelligence: From Natural to Artificial Systems.
Santa Fe Institute Studies in the Sciences of Complexity. Oxford
University Press, 198 Madison Avenue, New York, New York 10016,
United States of America.

Bibliography III



Bonabeau, E. and Théraulaz, G. (2000).

Swarm smarts (behavior of social insects as model for complex systems).

Scientific American, 282(3):72–79.



Bordini, R. H. and Hübner, J. F. (2006).

BDI agent programming in AgentSpeak using *Jason* (tutorial paper).

In Toni, F. and Torroni, P., editors, *Computational Logic in Multi-Agent Systems*, volume 3900 of *Lecture Notes in Computer Science*, pages 143–164. Springer.

6th International Workshop, CLIMA VI, London, UK, June 27–29, 2005. Revised Selected and Invited Papers.

Bibliography IV



Bordini, R. H., Hübner, J. F., and Wooldridge, M. J. (2007).
Programming Multi-Agent Systems in AgentSpeak using Jason.
John Wiley & Sons, Ltd.
Hardcover.



Bratman, M. E. (1987).
Intention, Plans, and Practical Reason.
Harvard University Press.



Bratman, M. E. (1990).
What is intention?
In Cohen, P. R., Morgan, J. L., and Pollack, M. E., editors, *Intentions
in Communication*, pages 15–32. The MIT Press, Cambridge, MA.

Bibliography V



Busi, N., Gorrieri, R., and Zavattaro, G. (1998).

A process algebraic view of Linda coordination primitives.

Theoretical Computer Science, 192(2):167–199.



Camazine, S., Deneubourg, J.-L., Franks, N. R., Sneyd, J., Theraulaz, G., and Bonabeau, E. (2001).

Self-Organization in Biological Systems.

Princeton Studies in Complexity. Princeton University Press, 41

William Street, Princeton, New Jersey 08540, United States of America.

Bibliography VI



Castelfranchi, C. (1995).

Guarantees for autonomy in cognitive agent architecture.

In Wooldridge, M. J. and Jennings, N. R., editors, *Intelligent Agents*, volume 890 of *Lecture Notes in Computer Science*, pages 56–70.

Springer Berlin Heidelberg.

ECAI-94 Workshop on Agent Theories, Architectures, and Languages (ATAL) Amsterdam, The Netherlands 8–9 August 1994. Proceedings.



Ciancarini, P. (1996).

Coordination models and languages as software integrators.

ACM Computing Surveys, 28(2):300–302.

Bibliography VII



Ciancarini, P., Omicini, A., and Zambonelli, F. (2000).
Multiagent system engineering: The coordination viewpoint.
In Jennings, N. R. and Lespérance, Y., editors, *Intelligent Agents VI. Agent Theories, Architectures, and Languages*, volume 1757 of *LNAI*, pages 250–259. Springer.
6th International Workshop (ATAL'99), Orlando, FL, USA, 15–17 July 1999. Proceedings.



Conte, R. and Castelfranchi, C., editors (1995).
Cognitive and Social Action.
Routledge.



De Nicola, R., Ferrari, G., and Pugliese, R. (1998).
KLAIM: A kernel language for agent interaction and mobility.
IEEE Transaction on Software Engineering, 24(5):315–330.

Bibliography VIII



De Nicola, R., Latella, D., Katoen, J.-P., and Massink, M. (2006).
StoKlaim: A stochastic extension of Klaim.
Technical Report 2006-TR-01, Istituto di Scienza e Tecnologie
dell'Informazione “Alessandro Faedo” (ISTI).



De Wolf, T. and Holvoet, T. (2005).
Emergence versus self-organisation: Different concepts but promising
when combined.
In Brueckner, S., Di Marzo Serugendo, G., Karageorgos, A., and
Nagpal, R., editors, *Engineering Self Organising Systems:
Methodologies and Applications*, volume 3464 of *LNCS (LNAI)*, pages
1–15. Springer.



Dennett, D. (1971).
Intentional systems.
Journal of Philosophy, 68:87–106.

Bibliography IX



Dennett, D. (2007).

Intentional systems theory.

In Beckermann, A. and Walter, S., editors, *Oxford Handbook of the Philosophy of Mind*, chapter 19. Oxford University Press.



Di Marzo Serugendo, G., Gleizes, M.-P., and Karageorgos, A. (2011).

Self-organising software.

In Di Marzo Serugendo, G., Gleizes, M.-P., and Karageorgos, A., editors, *Self-organising Software. From Natural to Artificial Adaptation*, Natural Computing Series, chapter 2, pages 7–32.

Springer Berlin Heidelberg, Berlin, Heidelberg.



Dorigo, M. and Stützle, T. (2004).

Ant Colony Optimization.

MIT Press, Cambridge, MA.

Bibliography X



Dorigo, M., Tuci, E., Mondada, F., Nolfi, S., Deneubourg, J.-L., Floreano, D., and Gambardella, L. M. (2005).

The SWARM-BOTS project.

Künstliche Intelligenz, 4/05:32–35.

Also available at <http://www.swarm-bots.org> as IRIDIA Technical Report No. TR/IRIDIA/2005-018.



Fensel, D. (2004).

Triple-space computing: Semantic web services based on persistent publication of information.

In Aagesen, F. A., Anutariya, C., and Wuwongse, V., editors, *Intelligence in Communication Systems*, volume 3283 of *LNCS*, pages 43–53.

IFIP International Conference (INTELLCOMM 2004), Bangkok, Thailand, 23–26 November 2004. Proceedings.

Bibliography XI



Freeman, E., Hupfer, S., and Arnold, K. (1999).
JavaSpaces Principles, Patterns, and Practice: Principles, Patterns and Practices.
The Jini Technology Series. Addison-Wesley Longman.



Gardelli, L., Viroli, M., Casadei, M., and Omicini, A. (2007).
Designing self-organising MAS environments: The collective sort case.
In Weyns, D., Parunak, H. V. D., and Michel, F., editors,
Environments for MultiAgent Systems III, volume 4389 of *LNAI*, pages
254–271. Springer.
3rd International Workshop (E4MAS 2006), Hakodate, Japan,
8 May 2006. Selected Revised and Invited Papers.

Bibliography XII



Gelernter, D. (1985).

Generative communication in Linda.

ACM Transactions on Programming Languages and Systems,
7(1):80–112.



Gelernter, D. and Carriero, N. (1992).

Coordination languages and their significance.

Communications of the ACM, 35(2):97–107.



Goldstein, J. (1999).

Emergence as a construct: History and issues.

Emergence, 1(1):49–72.

Bibliography XIII



González Pérez, P. P., Omicini, A., and Sbaraglia, M. (2013).
A biochemically-inspired coordination-based model for simulating
intracellular signalling pathways.

Journal of Simulation.

Special Issue on Agent-based Modeling and Simulation.



Graesser, S. F. A. (1996).

Is it an agent, or just a program?: A taxonomy for autonomous
agents.

In Müller, J. P., Wooldridge, M. J., and Jennings, N. R., editors,
*Intelligent Agents III Agent Theories, Architectures, and Languages:
ECAI'96 Workshop (ATAL) Budapest, Hungary, August 12–13, 1996
Proceedings*, volume 1193 of *Lecture Notes In Computer Science*,
pages 21–35. Springer.

Bibliography XIV

-  Grassé, P.-P. (1959).
La reconstruction du nid et les coordinations interindividuelles chez *Bellicositermes natalensis* et *Cubitermes* sp. la théorie de la stigmergie: Essai d'interprétation du comportement des termites constructeurs.
Insectes Sociaux, 6(1):41–80.
-  Kephart, J. O. and Chess, D. M. (2003).
The vision of autonomic computing.
Computer, 36(1):41–50.
-  Liu, J. and Tsui, K. C. (2006).
Toward nature-inspired computing.
Communications of the ACM, 49(10):59–64.

Bibliography XV



Mamei, M. and Zambonelli, F. (2004).

Programming pervasive and mobile computing applications with the TOTA middleware.

In *Pervasive Computing and Communications*, pages 263–273.

2nd IEEE Annual Conference (PerCom 2004), Orlando, FL, USA, 14–17 March 2004. Proceedings.



Mamei, M. and Zambonelli, F. (2006).

Field-Based Coordination for Pervasive Multiagent Systems. Models, Technologies, and Applications.

Springer Series in Agent Technology. Springer.

Bibliography XVI



Mariani, S. and Omicini, A. (2012a).

Molecules of Knowledge: Self-organisation in knowledge-intensive environments.

In Fortino, G., Bădică, C., Malgeri, M., and Unland, R., editors, *Intelligent Distributed Computing VI*, volume 446 of *Studies in Computational Intelligence*, pages 17–22. Springer.

6th International Symposium on Intelligent Distributed Computing (IDC 2012), Calabria, Italy, 24-26 September 2012. Proceedings.



Mariani, S. and Omicini, A. (2012b).

Self-organising news management: The *Molecules of Knowledge* approach.

In Fernandez-Marquez, J. L., Montagna, S., Omicini, A., and Zambonelli, F., editors, *1st International Workshop on Adaptive*

Bibliography XVII

Service Ecosystems: Natural and Socially Inspired Solutions (ASENSIS 2012), pages 11–16, SASO 2012, Lyon, France.
Pre-proceedings.



Mariani, S. and Omicini, A. (2013).

Event-driven programming for situated MAS with ReSpecT tuple centres.

In Klusch, M., Thimm, M., and Paprzycki, M., editors, *Multiagent System Technologies*, volume 8076 of *LNAI*, pages 306–319. Springer.
11th German Conference (MATES 2013), Koblenz, Germany, 16-20 September 2013. Proceedings.

Bibliography XVIII



McCarthy, J. (1979).

Ascribing mental qualities to machines.

In Ringle, M., editor, *Philosophical Perspectives in Artificial Intelligence*, Harvester Studies in Cognitive Science, pages 161–195. Harvester Press, Brighton.



Minsky, N. H. and Ungureanu, V. (2000).

Law-Governed interaction: A coordination and control mechanism for heterogeneous distributed systems.

ACM Transactions on Software Engineering and Methodology (TOSEM), 9(3):273–305.



Nardini, E., Omicini, A., and Viroli, M. (2013).

Semantic tuple centres.

Science of Computer Programming.

Special Issue on Self-Organizing Coordination. To appear.

Bibliography XIX

-  Nardini, E., Omicini, A., Viroli, M., and Schumacher, M. I. (2011). Coordinating e-health systems with TuCSon semantic tuple centres. *Applied Computing Review*, 11(2):43–52.
-  Odell, J. (2002). Objects and agents compared. *Journal of Object Technologies*, 1(1):41–53.
-  Omicini, A. (2013). Agents writing on walls: Cognitive stigmergy and beyond. In Paglieri, F., Tummolini, L., Falcone, R., and Miceli, M., editors, *The Goals of Cognition. Essays in Honor of Cristiano Castelfranchi*, chapter 29, pages 543–556. College Publications, London.

Bibliography XX



Omicini, A. and Denti, E. (2001).

From tuple spaces to tuple centres.

Science of Computer Programming, 41(3):277–294.



Omicini, A., Ricci, A., and Viroli, M. (2006).

The multidisciplinary patterns of interaction from sciences to Computer Science.

In Goldin, D. Q., Smolka, S. A., and Wegner, P., editors, *Interactive Computation: The New Paradigm*, pages 395–414. Springer.



Omicini, A., Ricci, A., Viroli, M., Castelfranchi, C., and Tummolini, L. (2004).

Coordination artifacts: Environment-based coordination for intelligent agents.

In Jennings, N. R., Sierra, C., Sonenberg, L., and Tambe, M., editors, *3rd international Joint Conference on Autonomous Agents and*

Bibliography XXI

Multiagent Systems (AAMAS 2004), volume 1, pages 286–293, New York, USA. ACM.



Omicini, A. and Zambonelli, F. (1999).

Coordination for Internet application development.

Autonomous Agents and Multi-Agent Systems, 2(3):251–269.

Special Issue: Coordination Mechanisms for Web Agents.



Parunak, H. V. D. (2006).

A survey of environments and mechanisms for human-human stigmergy.

In Weyns, D., Parunak, H. V. D., and Michel, F., editors, *Environments for Multi-Agent Systems II*, volume 3830 of *LNCS*, pages 163–186. Springer.

Bibliography XXII



Parunak, H. V. D., Brueckner, S., and Sauter, J. (2002).

Digital pheromone mechanisms for coordination of unmanned vehicles.

In Castelfranchi, C. and Johnson, W. L., editors, *1st International Joint Conference on Autonomous Agents and Multiagent systems*, volume 1, pages 449–450, New York, NY, USA. ACM.



Rao, A. S. and Georgeff, M. P. (1992).

An abstract architecture for rational agents.

In Nebel, B., Rich, C., and Swartout, W. R., editors, *3rd International Conference on Principles of Knowledge Representation and Reasoning (KR '92)*, pages 439–449, Cambridge, MA, USA. Morgan Kaufmann. Proceedings.

Bibliography XXIII



Rao, A. S. and Georgeff, M. P. (1995).

BDI agents: From theory to practice.

In Lesser, V. R. and Gasser, L., editors, *1st International Conference on Multi Agent Systems (ICMAS 1995)*, pages 312–319, San Francisco, CA, USA. The MIT Press.



Ricci, A., Omicini, A., Viroli, M., Gardelli, L., and Oliva, E. (2007).

Cognitive stigmergy: Towards a framework based on agents and artifacts.

In Weyns, D., Parunak, H. V. D., and Michel, F., editors, *Environments for MultiAgent Systems III*, volume 4389 of *LNCS*, pages 124–140. Springer.

3rd International Workshop (E4MAS 2006), Hakodate, Japan, 8 May 2006. Selected Revised and Invited Papers.

Bibliography XXIV



Ricci, A., Viroli, M., and Omicini, A. (2005).
Environment-based coordination through coordination artifacts.
In Weyns, D., Parunak, H. V. D., and Michel, F., editors,
Environments for Multi-Agent Systems, volume 3374 of *LNAI*, pages
190–214. Springer.
1st International Workshop (E4MAS 2004), New York, NY, USA,
19 July 2004. Revised Selected Papers.



Rossi, D., Cabri, G., and Denti, E. (2001).
Tuple-based technologies for coordination.
In Omicini, A., Zambonelli, F., Klusch, M., and Tolksdorf, R., editors,
*Coordination of Internet Agents: Models, Technologies, and
Applications*, chapter 4, pages 83–109. Springer.

Bibliography XXV



Sauter, J. A., Matthews, R. S., Parunak, H. V. D., and Brueckner, S. (2005).

Proceedings of the 4th international joint conference on autonomous agents and multiagent systems (aamas 2005).

In Dignum, F., Dignum, V., Koenig, S., Kraus, S., Singh, M. P., and Wooldridge, M., editors, *Proceedings of the 4th International Joint Conference on Autonomous Agents and Multiagent Systems (AAMAS 2005)*, pages 903–910, Utrecht, The Netherlands. ACM Press.



Simon, H. A. (1962).

The architecture of complexity.

Proceedings of the American Philosophical Society, 106(6):467–482.

Bibliography XXVI

 Solé, R. V. and Bascompte, J. (2006).
Self-Organization in Complex Ecosystems.

Number 42 in Monographs in population Biology. Princeton University Press, 41 William Street, Princeton, New Jersey 08540, United States of America.

 Tolksdorf, R. and Menezes, R. (2004).
Using Swarm Intelligence in Linda Systems.

In Omicini, A., Petta, P., and Pitt, J., editors, *Engineering Societies in the Agents World IV*, volume 3071 of *LNCS*, pages 49–65. Springer. 4th International Workshops (ESAW 2003), London, UK, 29-31 October 2003. Revised Selected and Invited Papers.

Bibliography XXVII



Viroli, M. and Casadei, M. (2009).

Biochemical tuple spaces for self-organising coordination.

In Field, J. and Vasconcelos, V. T., editors, *Coordination Languages and Models*, volume 5521 of *LNCS*, pages 143–162. Springer, Lisbon, Portugal.

11th International Conference (COORDINATION 2009), Lisbon, Portugal, June 2009. Proceedings.



Viroli, M., Casadei, M., Nardini, E., and Omicini, A. (2010).

Towards a chemical-inspired infrastructure for self-* pervasive applications.

In Weyns, D., Malek, S., de Lemos, R., and Andersson, J., editors, *Self-Organizing Architectures*, volume 6090 of *LNCS*, chapter 8, pages 152–176. Springer.

Bibliography XXVIII

1st International Workshop on Self-Organizing Architectures (SOAR 2009), Cambridge, UK, 14-17 September 2009, Revised Selected and Invited Papers.



Viroli, M. and Omicini, A. (2006).

Coordination as a service.

Fundamenta Informaticae, 73(4):507–534.

Special Issue: Best papers of FOCLASA 2002.



Viroli, M., Pianini, D., Montagna, S., and Stevenson, G. (2012).

Pervasive ecosystems: a coordination model based on semantic chemistry.

In Ossowski, S., Lecca, P., Hung, C.-C., and Hong, J., editors, *27th Annual ACM Symposium on Applied Computing (SAC 2012)*, Riva del Garda, TN, Italy. ACM.

Bibliography XXIX



Wegner, P. (1997).

Why interaction is more powerful than algorithms.

Communications of the ACM, 40(5):80–91.



Weyns, D., Schelfhout, K., Holvoet, T., and Lefever, T. (2005).

Proceedings of the 4th international joint conference on autonomous agents and multiagent systems (aamas 2005).

In Dignum, F., Dignum, V., Koenig, S., Kraus, S., Singh, M. P., and Wooldridge, M., editors, *Proceedings of the 4th International Joint Conference on Autonomous Agents and Multiagent Systems (AAMAS 2005)*, pages 67–74, Utrecht, The Netherlands. ACM Press.



Wooldridge, M. J. (2002).

An Introduction to MultiAgent Systems.

John Wiley & Sons Ltd., Chichester, UK.

Bibliography XXX



Zambonelli, F., Castelli, G., Ferrari, L., Mamei, M., Rosi, A., Di Marzo, G., Risoldi, M., Tchao, A.-E., Dobson, S., Stevenson, G., Ye, Y., Nardini, E., Omicini, A., Montagna, S., Viroli, M., Ferscha, A., Maschek, S., and Wally, B. (2011).

Self-aware pervasive service ecosystems.

Procedia Computer Science, 7:197–199.

Proceedings of the 2nd European Future Technologies Conference and Exhibition 2011 (FET 11).

Frontiers of Autonomous Systems

Andrea Omicini

`andreaomicini.apice.unibo.it`

Dipartimento di Informatica – Scienza e Ingegneria (DISI)
ALMA MATER STUDIORUM – Università di Bologna

Alma Graduate School
Bologna, Italy
May 8, 2014

